



Unpacking mystery boxes: the logical substructure of prima facie conception

Til Eyinck¹

Received: 4 July 2025 / Accepted: 9 July 2026
© The Author(s) 2026

Abstract

Prima facie conception can appear mysterious. Having a prima facie grasp of something means being ‘on the edge’ of conceiving. But it seems that often there is no way for us to know if what our intuitions *present as* prima facie conceivable is *actually* conceivable, i.e. conceivable under ideal rational reflection. We propose a novel account of prima facie conception: on the basis of working hypotheses concerning the logical interactions of conception, sensory imagination and prima facie conception, we isolate different subtypes of what we call *general prima facie conception*. With these tools, we explore (what seems to be) the deeper logical substructure of prima facie conception. We hint at how this could help us understand different reasons for different intuitions (for instance, regarding intuitive differences in the truth-value of conditionals with impossible antecedents).

Keywords Conceivability · Imagination · Impossibility · Modal epistemology · Intuitions

1 Introduction

We present: $\boxplus$, the mystery box. The mystery box is essentially a *metaphor* (for a function) that we will employ in this essay to discuss prima facie conception. Prima facie conception, as it is understood here, describes a ‘hazy’ state of mind involving something that we do not fully grasp; something we conceive of *at first glance*. Prima facie conception is closely related to epistemic uncertainty, since it arguably refers to the delicate interface between sensory data, the possibility of their cognition and genuine, propositional understanding.

✉ Til Eyinck
tileyinck@gmail.com

¹ Cologne Center for Contemporary Epistemology and the Kantian Tradition, University of Cologne, Albertus Magnus Platz, Cologne 50931, NRW, Germany

Philosophical thought experiments, for instance, often arguably arise from, or consist in, acts of *prima facie* conception; we usually start with under-determined scenarios that we ‘roughly assemble before our inner eye’, or even visualize in the form of detailed, yet at first incomplete, visual representations in our minds (cf. Stuart, 2022; Chalmers, 2002). Only then, *by reasoning*, do we examine the scenario’s coherence and can thus *sometimes* evaluate whether the respective scenario (or what we believed to be the scenario) is not only *prima facie*, but ‘truly’ conceivable, i.e. conceivable under ideal rational reflection (cf. Chalmers, 2002).¹ When mathematicians investigate whether an *unproven conjecture* that they *assume to be* a theorem is *de facto* a theorem, *prima facie* conception also seems to play a role: intuitions related to what they *prima facie conceive to be* a theorem seem to have some kind of justificatory force and may serve as a kind of guide for their inquiry (cf. Priest, 2016; Berto et al., 2018). Concerning the notion of *prima facie* conception that we employ here, we are almost completely in line with Chalmers, who assumes that:

S is *prima facie* conceivable for a subject when S is conceivable for that subject on first appearances. [...] S is ideally conceivable when S is conceivable on ideal rational reflection. It sometimes happens that S is *prima facie* conceivable to a subject, but that this *prima facie* conceivability is undermined by further reflection showing that the tests that are criterial for conceivability are not in fact passed. In this case, S is not ideally conceivable.² (Chalmers, 2002, p. 147)

Prima facie conception (as it is understood here) is therefore no guarantee for the conceivability of what we conceive (or intuit to conceive) *prima facie*; *prima facie* conception is in this respect *defeasible*, i.e. ‘half-sidedly mysterious’. It is taken to

¹ Ideal rational reflection is not easy to define, it could be understood in terms of a dynamic epistemic modal logic such as the one introduced by (Duc, 1997).

² According to Chalmers, and as we understand it here, ideal conception is thus how ideal, Godlike agents conceive, whereas *prima facie* conception is our ‘messy’ seeming to conceive. An objection might arise here: Shouldn’t an ideal being, *not* bound by computational resources or time, be able to conceive *just as* finite agents conceive? Shouldn’t such a being be able to conceive *even more*? There are two things to point out: (1) the theory presented below is compatible with the idea that an ideal being can conceive exactly what we can conceive of and even more; Suppose a human *believes* to be conceiving what is allegedly expressed by ‘x’, but in reality, without knowing that what ‘x’ expresses cannot be conceived, conceived something else, y. According to the approach presented here, an ideal being thus *can* conceive y without any problem (given that y entails no contradiction), just as we do. The only, but crucial, difference between such a being and us, according to the views developed here, is that the respective being always conceives at relevant depth (y in this case), while we don’t; cf. also DUC’s epistemic modal logic quoted in the previous footnote. In a nutshell, DUC’s idea is that ideal agents can conceive what we do, but cannot be in a state of insufficient computational depth. To anyone who is still dissatisfied, one could point out the following: (2) there are many things that Godlike, ideal beings, if they were to exist, arguably would not be able to do, such as lying. Cf., for instance, the views of Anselm of Canterbury on these issues; “If God is perfectly just, he cannot lie. But if God is omnipotent, how can there be something he cannot do? Anselm’s solution is to explain that omnipotence does not mean the ability to do everything; instead, it means the possession of unlimited power. Now the so-called “ability” or “power” to lie is not really a power at all; it is a kind of weakness. Being omnipotent, God has no weakness. So it turns out that omnipotence actually entails the inability to lie.” (Williams, 2025, ch. 3.2). There is, however, also a third dimension related to these worries that we address in the last section of the present work, related to computational depth and to the transparency of introspection of ideal and non-ideal agents.

refer to something that characterises specific epistemic states of finite, non-idealised agents who do not have infinite resources and time at their disposal. For an ideal being either conceives of something or does not, and if it does, it invariably does so with adequate computational depth. It is therefore in our view not productive to speak of ‘prima facie conception’ in the case of idealised reasoners. We, as finite agents, on the contrary, rather frequently imagine something *prima facie* that we do not understand at a relevant depth.³

‘Our approach differs, however, from Chalmers’ in that we do not ‘look’ directly ‘at’ *S*; instead of directly referring to what presumably is the content *S* of a *prima facie* conception, we make use of the mystery box $\boxplus$, as a placeholder (with a function). For now, you can think of the mystery box in terms of an urn that *might* come with nebulous beliefs about *what you hold to be* its content and sometimes with intuitions⁴ concerning what *prima facie appears to be* its content. The further reasons for this will hopefully become clear below.

Under the dogma that only non-contradictory sets of *propositions* and logically possible objects – like triangles – are conceivable/imaginable (a thesis that we do not want to defend here, but employ mainly as a working hypothesis), in this essay we try to investigate the logical substructure of *prima facie* conception. This means we try to ‘look in’ $\boxplus$. The dogma that only non-contradictory sets of propositions (and possible objects) are conceivable is what we call our ‘path’ in this essay, as this assumption assures a stable epistemological ‘top-structure’ against which we can explore the ‘substructure’ of *prima facie* conception.

This paper is organised as follows: In a first step, we collect and formalise some working hypotheses concerning plausible interactions of conception, sensory imagination and *prima facie* conception (Sect. 2). From these working hypotheses, it emerges that the logical substructure of *prima facie* conception must be ramified in a specific way. We trace these ramifications and thereby determine what we consider to be different subtypes of *prima facie* conception (Sect. 3). But we do not stop there; there is more substructure to unpack (Sect. 4). We then interpret the insights we have gained and hint at how they might help us to better understand the object and ground of, what seem to be, different reasons for different intuitions (Sect. 5). We then sum up (Sect. 6).

³ In other words, to anticipate this, we will argue that *prima facie* conception for ideal reasoners coincides with ideal propositional understanding. There is only conception for ideal reasoners, and partial conception/grasp of a proposition, i.e. what we will call *safe prima facie* conception, is limited to non-ideal, fallible agents. Regarding the modelling of cognitive agents in this vein, cf. both (Chalmers, 2002) and (Duc, 1997).

⁴ There are various ways of talking about intuition in philosophy. For instance, one might refer to an immediate, propositional, direct insight into ideas. Here, however, we use the term to refer to phenomena such as when a mathematician has the intuition that a particular method of proof might be correct, or when we have a ‘gut feeling’ about a particular *prima facie* conception (Antognazza & Segala, 2023). See also the final section of this present essay.

2 The top-structure

In order to discuss *prima facie* conception below, we will now first outline some plausible assumptions about imagination understood in a general sense:⁵

Imagination in general, arguably, can be broken down into (at least) two *modes* of imagination; sensory imagination (I), which involves the reproduction of sensory data, as in the case of imagistic mental content, and conceptual imagination (C) – sometimes also called attitudinal imagination –, which involves some level of conceptual understanding but does not necessarily involve the reproduction of sensory data. For example, if one assumes that St Andrews is not in Scotland but further south, this does not necessarily seem to go hand in hand with *vivid mental images* (i.e. with visual representation or simulations of the world) in our minds (cf. Kind, 2023).

Another example, which comes from Berto, is possibly even clearer; what does it mean to conceive of the fact that stock prices are falling worldwide? It seems that no particular visual mental image can be identified here that corresponds to the content of the entirely possible conception (cf. Berto et al., 2022, p. 120).⁶

Following Descartes, our *introspective experience* when we try to conceive of a chiliagon (a geometrical figure with one thousand corners and therefore sides) seems to suggest that we can conceive of logically possible objects (such as the chiliagon) that are not necessarily imaginable in the sense of visual representations in our minds (cf. Descartes, 1996). So, even if it may be that for *idealised* agents sensory imagination and conception *could* conflate (which we do not assume here⁷), for finite agents sensory imagination arguably always implies an act of conception, but an act of conceptual imagination should not be understood as necessarily implying sensory imagination. It is, thus, arguably fair to assume that, whenever finite cognitive agents ‘bring something to their mind’, ‘understand something’, etc., this involves an act of conceiving (C) but not always an act of (sensory) imagination (I).⁸

⁵These assumptions, for the most part, are considered to be more or less established assumptions (Gendler & Hawthorne, 2002; Chalmers, 2002).

⁶We are aware that these assumptions correspond to an oversimplification of the positions in the literature concerning different types of imagination. However, we decide to keep these matters simplified in order to avoid complications that could arise on the formal side when trying to capture more fine-grained notions of imagination/conception. For our aim is not, strictly speaking, to make a contribution to the philosophy of imagination, but to examine *prima facie* conception.

⁷According to Brown, we are in line not only with him, but with Etchemendy, Barwise and Wittgenstein concerning this dichotomy (cf. Brown, 1997): “It is probably true that anything can stand for anything. But it is not true that anything can stand *pictorially* for anything. Something special is needed. But what is it about a picture of X that makes it a *picture* of X? [...] In the *Tractatus*, Wittgenstein made a few cryptic remarks about the relation between pictures and what they picture. [...] [*Tractatus*, 2.161]; [...] [*Tractatus*, 2.18]. What this suggests is a kind of structural similarity, a notion which is captured by the concept of an isomorphism. Barwise and Etchemendy (who are among the very few sympathetic to the use of pictures in inference) explicitly adopt such a view. They hold that ‘a good diagram is isomorphic, or at least homomorphic, to the situation it represents’.”

⁸Indeed, one could argue at this point that one can very well imagine things without having a concept of them, that therefore an act of visual imagination does not necessarily have to imply an act of conceptual understanding. Think of phenomena such as freedom or love. We take note of this objection, but have no space to discuss why we believe this stance to be false.

Furthermore, we assume that imagination in general (comprising both C and I, the two ways to grasp a proposition) encompasses not only those mental processes which, to borrow Wansing's term, are subject to "doxastic control" (Wansing, 2017, p. 2844). In order to clarify what is meant, think of it this way: if someone addresses you from the shadows, without being seen by you beforehand, and you understand what he or she is saying to you, this may trigger visual images in your mind's eye that are not voluntary, but which fall under the concept of imagination employed here. Or think of involuntary imagination under the influence of drugs, for instance. The following discussion is also concerned with such possible imaginative content.⁹

To further discuss the topics at hand from now on, we introduce the following notions:¹⁰ I (imagination), C (conception)¹¹, PFC (prima facie conception), $\Diamond$ (logical possibility). We also introduce the interactions of $\Diamond$, I and C: $\Diamond C$ (conceivability) and $\Diamond I$ (imaginability). With the exception of PFC¹², these are *all* taken to be normal, diamond-like *modal operators* (for we take them to quantify over possible worlds only¹³).

Costa-Leite (2010) showed how the above notion of interaction of conceptual and sensory imagination can be captured by what he, on the basis of their respective philosophical provenance, calls *The Law of Descartes-Vasilev* and *The First Law of Hume* (cf. also Gendler & Hawthorne, 2002, pp. 13–26):

- (1) $I\varphi \rightarrow C\varphi$ *Descartes-Vasilev Law* (imagination implies conception),
- (2) $C\varphi \rightarrow \Diamond\varphi$ *The First Law of Hume* (conception implies possibility).

From 1 and 2 we also get *The Second Law of Hume*: $I\varphi \rightarrow \Diamond\varphi$ (imagination implies possibility). 1 and 2 are thus sufficient to define the nesting of conception, (sensory) imagination and logical possibility that is represented by Fig. 1 below.

If one excludes contradictory sets of propositions being conceivable from the outset, at first glance, it could seem plausible that even the most idealised agent can *at most* conceive of what corresponds to the set of all logically possible worlds (W) and of all possible objects (like triangles and so on). *In view of our path*, that is if one excludes the conceivability of contradictory matters, it is, however, arguably mistaken to assume that even the most ideal agent can conceive of the totality of possible worlds (W). To clarify what this stance relies on, assume W is given exten-

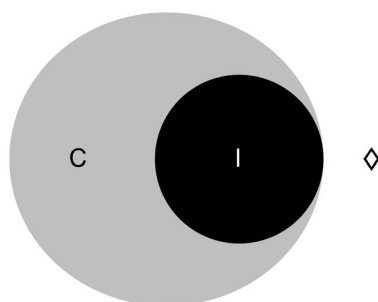
⁹ (It should be briefly noted that examples such as these can, of course, also be constructed using merely conceptual, i.e. non-sensory/imagistic propositional content, as in the case of the stock-market example. The aim here is simply to draw attention to doxastic control).

¹⁰ The idea of a modal interaction logic of this kind goes back to Costa-Leite, who developed IMAG. Formally speaking, this paper is based on the idea of extending IMAG by prima facie operators (see below in the text). Since it will not always be possible to specify everything that goes back to Costa-Leite's work, it should be emphasized here how important it is for this paper. Cf. Costa-Leite (2010).

¹¹ For the sake of economy and to avoid any confusion arising from similar terminology, we will often just use the terms 'imagination' and 'conception' in what follows, although, as mentioned above, we are actually referring to *sensory imagination* and *conceptual imagination* with the two.

¹² This *apparent* 'operator' is trickier to define, which is the whole point of this essay.

¹³ This means that, given that $\#$ is the respective operator, the semantics for the respective operator are: $\#\phi$ is true at a world w iff $\exists v(wRv \wedge \phi$ is TRUE at v).

Fig. 1 Nesting model for imagination and conception

sionally: in an extensional account of logical space, given two atomic propositions φ and ψ , the set of all (logically) possible worlds amounts to four possible worlds: $W\{\varphi \wedge \psi, \neg\varphi \wedge \psi, \varphi \wedge \neg\psi, \neg\varphi \wedge \neg\psi\}$.

Assuming that an idealised agent can conceive of the totality of what is logically possible (W), however, would already imply logical contradictions to be conceivable: for what could appear to be the conception of all logically possible worlds, in our example $C[(\varphi \wedge \psi) \wedge (\neg\varphi \wedge \psi) \wedge (\varphi \wedge \neg\psi) \wedge (\neg\varphi \wedge \neg\psi)]$, is, of course, contradictory. This just means that according to this notion of conceivability you might be able to conceive that *you are in the library and reading* ($\varphi \wedge \psi$) and also that *you are not in the library and reading* ($\neg\varphi \wedge \psi$), but that you cannot conceive of the scenario that *you are at the library and reading and also not reading* $[(\varphi \wedge \psi) \wedge (\varphi \wedge \neg\psi)]$. Since we are exploring the ‘path’ that has it that impossibilities are inconceivable, we should not understand W as a candidate for the maximal propositional content conceivable by an idealised agent (i.e. as what we want conceivability, $\Diamond C$, to express). Since the ‘path’ demands that we exclude contradictory ‘chunks’ of propositions to be conceivable ($\Diamond C$), we deny that W is ideally conceivable and we also exclude any contradictory ‘chunk’, i.e. any contradictory subsets of possible worlds like $(\varphi \wedge \psi) \wedge (\neg\varphi \wedge \psi)$, to be ideally conceivable too.

What has just been said is, in a way,¹⁴ why Salmon emphasized that one should acknowledge that the notion of *impossible* worlds was already contained in Kripke’s original idea of *possible* worlds:

Although it has not been generally recognized, impossible worlds already arise in classical Kripke-style possible-world semantics for propositional modal logic, in its T and B models. If w_2 is not accessible to w_1 , and w_1 is the actual world, then w_2 is an impossible world. (Salmon, 2024, p. 854)

We note, however, that free worlds as understood by Salmon (and possibly Kripke) are not necessarily impossible due to their inconsistency with other worlds, but merely due to a lack of accessibility. We will make use of the ‘impossible chunks’-interpretation of possible world semantics: However, importantly, *not* by ‘reinterpreting’ $\Diamond C$ as quantifying over contradictory ‘chunks’ of W and thereby implying that

¹⁴ Salmon is concerned with Kripke models in which worlds are not accessible because there are no corresponding condition on frames (see below in the main text), but such free worlds, arguably, *can sometimes be interpreted as* contradictory chunks of the extensional logical space.

ideal agents could conceive of logical contradictions arising from ‘free’, i.e. inaccessible worlds in logical space.

In the account proposed here, ideally conceivable ($\Diamond$ C) shall *maximally* apply to non-contradictory sets of propositions and possible objects – be W determined extensionally or not. Ideally conceiving a set of propositions here is thus taken to imply that the conceived set of propositions is *free from contradictions*. We can capture this as a general interaction principle with what we call *the law of ideal conceivability*. The law of ideal conceivability says that any ideally conceivable ($\Diamond$ C) *set of propositions* (Δ) is necessarily non-contradictory ($\Delta \not\vdash \perp$) and vice versa:

$\Diamond C\Delta \longleftrightarrow (\Delta \not\vdash \perp)$ (the dogmatic ‘path’).

(The mere possibility of forming contradictory chunks of per se conceivable propositions, which, going with Salmon, is not excluded by possible world frameworks from the outset, here, however, shall not guarantee for the conceivability of a corresponding contradiction of two per se conceivable conjuncts).

Let us now, to go on, also formally capture the already discussed hypothesis that *some* contents of conceptions are sensorially imaginable, but that not every *conceptual* imagination (C) goes hand in hand with a *sensory* imagination (I). Any conceivable set of propositions Δ is either imaginable ($\Diamond$ I) or not:

$C\Delta \rightarrow [(\Diamond I\Delta) \vee \neg(\Diamond I\Delta)]$ (conception implies possible imaginability).

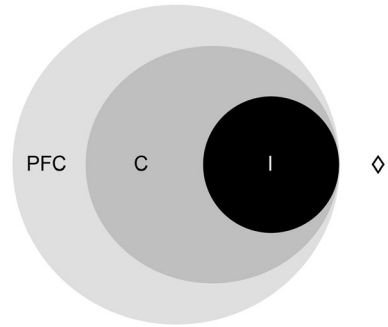
What else is there to say about I and C? We argue that it is plausible that the contents of two different acts of imaginations/conceptions (two possible propositional contents), should obey a principle of distribution:¹⁵ If an idealized agent imagines $\varphi \wedge \psi$, the agent must arguably be able to imagine φ ‘separately’ from ψ , i.e. the ideal agent must be able to imagine φ and also able to imagine ψ . The same is arguably desirable for conceiving. But one must be careful when formalising this, since one wants to avoid that contradictions arise in case, say, ψ implies $\neg\varphi$; in this case, it is reasonable to assume that an ideal agent might be able to imagine ψ and might be able to imagine φ , but is nevertheless not able to imagine $\psi \wedge \varphi$, since this would imply the imaginability of a contradiction (and we assumed this to be impossible). These desiderata for distributivity and mergeability should thus arguably be captured by the following principles that incorporate the conceivability condition ($\Diamond C\Delta \longleftrightarrow (\Delta \not\vdash \perp)$):

$[(\varphi \wedge \psi) \not\vdash \perp] \rightarrow [(\Diamond I\varphi \wedge \Diamond I\psi) \longleftrightarrow \Diamond I(\varphi \wedge \psi)]$ (imagination distribution),
 $[(\varphi \wedge \psi) \not\vdash \perp] \rightarrow [(\Diamond C\varphi \wedge \Diamond C\psi) \longleftrightarrow \Diamond C(\varphi \wedge \psi)]$ (conception distribution).¹⁶

¹⁵ Cf. the similar principles of distribution that are *validated* by IMAG, a language based on the fusion of the interaction-axioms (Descartes-Vasilev’s and Humes’s Laws) and three semantically identical (they all have ‘diamond-like’ truth conditions) normal modal logics (for possibility, conception and imagination): Costa-Leite (2010).

¹⁶ Note that one can conceive of a non-extended colour and of a non-coloured extension; but this does not mean that one can therefore also have an imagistic representation of a non-extended colour. However, since Descartes-Vasilev merely posits a dependency in the other direction (imagination implies conception), this arguably provides no counterexample to the principles. Nevertheless, we observe that the merging/distribution in the case of mere conception is perhaps less controversial.

Fig. 2 Nesting of imagination, conception and prima facie conception



These laws say that distribution (and merging) for acts of imagination and conception are possible for ideal agents as long as the conjunction of the respective propositions does not imply a contradiction.

Now one might think that the phenomenon of prima facie, as we have roughly outlined it so far according to Chalmers, can be captured in analogy to how conception relates to visual imaginability ($C\Delta \rightarrow [(\Diamond I\Delta) \vee \neg(\Diamond I\Delta)]$). This means one might think of prima facie conception assuming that imagination and conception always *simply* prima facie conception, but that, ‘looking’ from the ‘level below’, PFC either leads to conceivability or not (and that *if* it leads to conceivability, then possibly to imaginability):

$C\varphi \rightarrow PFC\varphi$ (conception implies prima facie conception),

$I\varphi \rightarrow PFC\varphi$ (imagination implies prima facie conception),

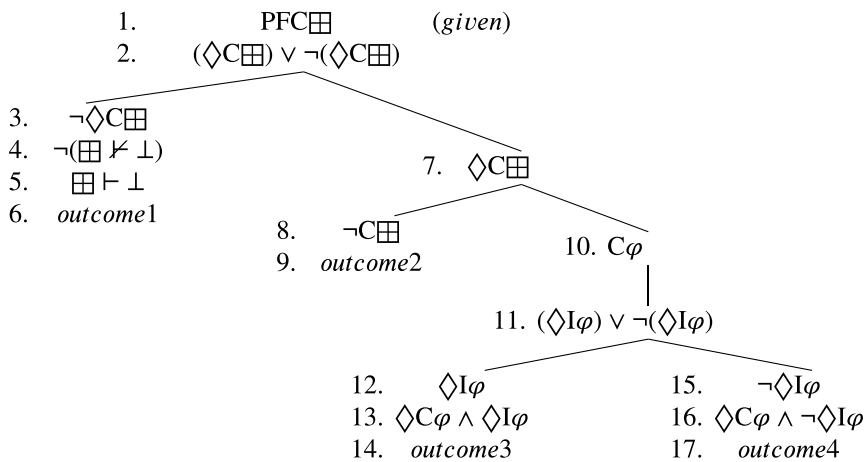
$PFC\Box \rightarrow [(\Diamond C\Box) \vee \neg(\Diamond C\Box)]$ (PFC implies possible conceivability).

These assumptions about prima facie conception result in the nesting represented in Fig. 2. We call this the *naïve view of prima facie conception*. We term it the naïve view because, as we now show, this view is unsatisfactory.

We now show the extent to which PFC understood in this naïve way is unsatisfactory by first showing the implications of our assumptions so far. Here are all relevant previous assumptions:

- (1) $I\varphi \rightarrow C\varphi$ (*Descartes-Vasilev Law*: imagination implies conception)
- (2) $C\varphi \rightarrow \Diamond\varphi$ (*The First Law of Hume*: conception implies possibility)
- (3) $I\varphi \rightarrow PFC\varphi$ (imagination implies prima facie conception)
- (4) $\Diamond C\Delta \longleftrightarrow (\Delta \not\vdash \perp)$ (the dogmatic ‘path’)
- (5) $C\Delta \rightarrow [(\Diamond I\Delta) \vee \neg(\Diamond I\Delta)]$ (conception implies possible imaginability)
- (6) $PFC\Box \rightarrow [(\Diamond C\Box) \vee \neg(\Diamond C\Box)]$ (PFC implies possible conceivability)

Let us now assume a (naïve) prima facie act of conception ($PFC\Box$) of a rational agent. Assume that *something* is ‘in’ $\Box$; assume *something* is prima facie conceived. Assuming $PFC\Box$ and the above, results in what follows:



Let us interpret this intuitive, tableaux-like representation. We draw attention to all four possible ‘outcomes’ of prima facie conception as it is understood so far and also to the branch 7–10, i.e. the case of a prima facie conception where the content of the mystery box ‘comes to light’ in the form of a conceivable proposition/possible object:

Outcome 1 can be interpreted as the result of a case of prima facie conception where one observes the non-conceivability of the respective contents of the mystery box. In this scenario, it seems nothing meaningful can be said about the content of the mystery box (for it is inconceivable according to our axioms). Outcome 2 can be interpreted as the case in which the contents of the mystery box are ideally conceivable, but are not actually conceived of.

The transition from 7 to 10 can be interpreted in terms of representing *the safe pathway to conceivability* that prima facie conceivability *can* be (prima facie conception implies possible conceivability and imaginability): Outcome 3 and 4 describe these uncontroversial ‘results’ of an act of prima facie conception arising from our axioms, namely the possibility that, generally speaking, the content of an act of prima facie conception *can* imply the conceivability and imaginability of what is prima facie conceived. In this case, metaphorically speaking, we are ‘allowed’ to open the box, since $I\varphi \rightarrow PFC\varphi$ and $C\varphi \rightarrow PFC\varphi$. For these two scenarios, i.e. *possible* outcomes of some acts of prima facie conception (outcome 3 and 4), it seems we can state that $PFC\varphi$.

Now it is hopefully becoming clear why we introduced the metaphor of the mystery box $\Box$ from the beginning on: given our ‘path’, there seems to be no way to know what agents prima facie conceive of in the case of outcome 1. And if one therefore, in such cases of prima facie conception, cannot know if what an agent (allegedly) prima facie conceived is even conceivable, it follows, in view of our path, that *we cannot be sure that there is any propositional content* (or possible object) in the box to be identified. Metaphorically speaking, we will try to shed some more light on this in this essay (we try to ‘illuminate’ the obscure parts of the inside of the mystery box). First, however, we must discuss why the prima facie account provided so far is unhelpful.

2.1 Why naive *prima facie* is *unhelpful*

Assume p is a theorem of geometry, i.e. a mathematical conjecture that has proven to be true. Assume further that an agent c is an excellent mathematician. The agent c can not only provide an analytical proof of what she at first only *prima facie* conceived of (PFC p), but can also *pictorially imagine* a proof of p in her mind, as many mathematicians are able to do:¹⁷ Picture proofs [in mathematics] are obviously too effective to be dismissed and they are potentially too powerful to be ignored. Making sympathetic sense of them is what is required of us. (Brown, 1997, p. 179)

The mathematician Felix Klein wrote about *imaginative, visual/pictorial* reasoning in mathematics:

Mathematics is not merely a matter of understanding but quite essentially a matter of imagination. (Arana, 2016, p. 463)

According to these quoted thinkers, the excellent mathematician c of our example would arguably fulfil Ip, Cp, and PFC p (she would *imagine* p , *conceive of* p and, since this is the most basic way of conceiving, *prima facie conceive of* p). Now suppose other, less talented mathematicians also tried to work on the theorem p , before it was proven to be a theorem by c .

Priest (2016) and other impossibilists like Jago and Berto (Cf. Berto & Jago, 2019, p. 21 ff.; Berto, 2017) would argue at this point that the above notion of *prima facie* conceivability is *helpful* in order to describe the successful mathematician, but *unhelpful* when it comes to describe what goes on in the mind of all those ‘unsuccessful’ mathematicians that also (according to impossibilists) worked on the same problem but held the negation of p to be true. These mathematicians (so they would argue) must have *prima facie* conceived of the negation of p (PFC $\neg p$):

Some things that are epistemically possible would seem to be logically impossible. Thus, before Wiles’ proof of the truth of Fermat’s last theorem, its negation was epistemically possible, though logically impossible. (Priest, 2016, p. 2652)

Let us turn back to our example of the geometrical theorem p . Priest, as we have just seen, apparently would argue that the ‘unsuccessful mathematicians’ must have conceived of *something*, when they *prima facie* conceived of $\neg p$. Priest proposes to

¹⁷ Besides the references following down below in the text, cf. also, classically, the introduction of *Geometry and the Imagination* by David Hilbert and Stephan Cohn-Vossen (1952) with regard to the dichotomy of conception and *visual* imagination (Hilbert & Cohn-Vossen, 1952). Note furthermore, how Brown, concerning a specific case, ascribes a kind of *prima facie* reliability to the respective picture proof that is superior to what is available in terms of an alternative analytic proof: “One of these (the picture) is *prima facie* reliable. The other (the analytic proof) is questionable, but our confidence in it is greatly enhanced by the fact that it agrees with the reliable method.” Brown (1997, 166).

broaden our notion of conception (and likely would suggest this also for *prima facie* conception), assuming that $PFC \neg p$ and even $C \neg p$ and potentially $I \neg p$ are possible epistemic acts (p still being the imaginable *theorem* of geometry, i.e. something true in every world where the laws of logic hold, i.e. true in every world in W). Therefore, Priest, like other impossibilists,¹⁸ suggests to give up the idea that only non-contradictory matters can be conceived of ($\Diamond C \Delta \longleftrightarrow (\Delta \not\vdash \perp)$). Priest and other impossibilists suggest adapting formal semantics based on intuitions (intuitions related to what is *inconceivable* according to ‘our path’) and they propose to ‘expand’ the realm of what is absolutely conceivable to include the realm of logical impossibility – that is, to assume that the realm of what is logically impossible has always been an integral part of it.¹⁹

We do agree with Priest and other impossibilists on the fact that a possibilist notion of conception (inherent to the notion of *prima facie* conception we have discussed so far) is unhelpful in describing what goes on in the mind of the ‘unsuccessful’ mathematicians (and possibly in all minds that *appear* to conceive of logical impossibilities). But we explore the opposite path when it comes to the consequences this should have for our logical theory: Instead of adapting formal semantics to our intuitions, we explore the logical substructure of *prima facie* conception. *We keep the above assumptions about what is ideally conceivable as they are.* We instead inquire what arguably *must* be ‘in’ $\boxplus$, given our ‘path’, by providing in what follows a slightly more fine-grained, recursive account of *prima facie* conception.

3 Logical substructure(s) of *prima facie* conception

We now begin to explore the logical substructure of *prima facie* conception. For this purpose, let us first assume that both outcomes of the ‘visible’ top-structure (3 and 4 in the tableaux) can each be matched to an underlying, different *type* of *prima facie* conception; In this sense, we first distinguish the above observed two subtypes of *prima facie* conception that are, ideally speaking, safe guides to possibility:

¹⁸ Cf (Priest, 2016; Berto, 2017, 2014; Berto & Jago, 2019; Tanaka, 2018).

¹⁹ One possible strategy for formally realizing the desired, expanded conception of conceivability involves introducing logically *impossible* worlds. Impossible worlds are, roughly put, logical safe-spaces within a given language L where fundamental logical principles, like the law of excluded middle, are ‘allowed’ to fail so that, in a way, the negation of theorems is ‘allowed’ to hold (cf. Berto et al., 2018). We would like to note that, under a particular interpretation, even proponents of a logically conservative modal epistemology such as Williamson can be read as impossibilists. For instance, the way he discusses epistemic processes that lead us to become aware of metaphysical necessities seems to imply that we do not merely conceive of what is logically possible: We determine that A is necessary by imagining the negation of A and counterfactually deducing an inconsistency: if $\neg A$ were true, a contradiction would follow. Because a contradiction is logically impossible, $\neg A$ cannot be true. It follows, therefore, that A must be true, rendering it necessary. Throughout this process, we conceive of $\neg A$ as that which is ultimately revealed to be an absolute impossibility. However, there is also an interpretation of Williamson that suggests that in such cases there is suppositional reasoning *without* conception at work, thereby placing him on the side of the possibilists. Cf. Williamson (2007), ch. 5).

- (1) **PFCI** (safe prima facie imagination that ideally implies imagination and conception),
- (2) **PFCC** (safe prima facie conception that ideally implies conception but does never imply imagination).

We assume that these subtypes of prima facie, the safe subtypes, are those safe ‘paths’ to conceivability and logical possibility we observed as interpretable scenarios on the top-structure for prima facie acts (hence ‘safe’). These subtypes of prima facie arguably can be captured as follows:

- $I\varphi \rightarrow \mathbf{PFCI}\varphi$ (imagination implies *safe prima facie imagination*)
- $\mathbf{PFCI}\varphi \rightarrow \Diamond I\varphi$ (*safe prima facie imagination* implies imaginability and therewith conceivability)
- $C\varphi \rightarrow \mathbf{PFCC}\varphi$ (conception implies *safe prima facie conception*)
- $\mathbf{PFCC}\varphi \rightarrow \Diamond C\varphi$ (*safe prima facie conception* implies only conceivability)

We therefore demystify $\boxplus$ with respect these two possible, safe cases of prima facie conception. For we assume, as our axioms suggest, that in these two cases, we are ‘allowed’ to ‘look’ into the box; we take these to guarantee for conceivability/imaginability, that is guaranteed possibility of conception/imagination under ideal rational reflection. This means that ideal conceivability/imaginability of φ arguably allows us in turn to assume an *idealised* act of conception of φ that relates to the respective prima facie conception of φ and thus ‘sheds light’ into the respective mystery box.

For these safe and transparent types of prima facie conception it furthermore seems plausible to assume a principle of distributivity and mergeability, just as we did for conception and imagination:

$$[(\varphi \wedge \psi) \nVdash \perp] \rightarrow [(\Diamond \mathbf{PFCC}\varphi \wedge \Diamond \mathbf{PFCC}\psi) \longleftrightarrow \Diamond \mathbf{PFCC}(\varphi \wedge \psi)],$$

(safe prima facie conception distribution and merging),

$$[(\varphi \wedge \psi) \nVdash \perp] \rightarrow [(\Diamond \mathbf{PFCI}\varphi \wedge \Diamond \mathbf{PFCI}\psi) \longleftrightarrow \Diamond \mathbf{PFCI}(\varphi \wedge \psi)],$$

(safe prima facie imagination distribution and merging).

We are approaching the core idea of this essay that consists in a structural argument; for given these two safe subtypes, one might argue that, *symmetrically* to the above assumptions about safe types, there must be two other prima facie types that cannot go hand in hand with imaginability and conceivability (let us call them ‘dead-end’-types of prima facie conception). For without assuming prima facie dead-ends, the whole concept of prima facie would be reduced to absurdity, since prima facie, in light of our path, was only introduced in order to ‘separate the wheat from the chaff’, i.e. in order to separate what, ideally speaking, guides to conceivability/imaginability from what does *never*:

- (1) **PFCC** (*dead-end* prima facie conception),
- (2) **PFCI** (*dead-end* prima facie imagination).

These seem to be the *obscure* types of prima facie conception that arguably create problems ‘in the top-structure’ – they are part of the reason for the mystery box metaphor. For if the two above safe cases of prima facie conception (**PFCC**, **PFCI**) are ‘made transparent’, there must be another ‘obscure, non-transparent’ subtype of prima facie conception/imagination to be held accountable for the unsafety of prima facie conception. But how could one capture this, formally speaking? And what is *the content* of these acts of prima facie conception, if it is not *not* a conceivable proposition or a possible object?

In Stalnaker’s tradition,²⁰ we argue that these types of prima facie conception could be analysed as being formulas that map *sentences* to *propositions*. This means we define dead-end prima facie conception/imagination by making use of interactions of prima facie types: in the case of *dead-end* prima facie imagination, we assume that whenever cognitive agents prima facie imagine in this way, the respective agents are *safely* prima facie imagining a *proposition*, something conceivable, ideally speaking. (We do not yet discuss the obvious follow-up question of how the function between sentences and propositions could be thought of; we will get back to this soon). We thus propose the following: In view of a **PFCI**-type act of prima facie involving a *sentence* like ‘*colourless squared circles are green*’²¹, one can interpret **PFCI** *colourless squared circles are green* as:

- **PFCI***colourless squared circles are green* → **PFCI** φ (dead-end prima facie imagination implies safe prima facie imagination).

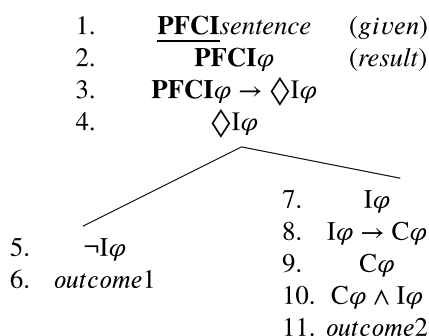
We assume this to be the case also for dead-end prima facie *conception*:

- **PFCC***sentence* → **PFCC** φ (dead-end prima facie conception implies safe prima facie conception).

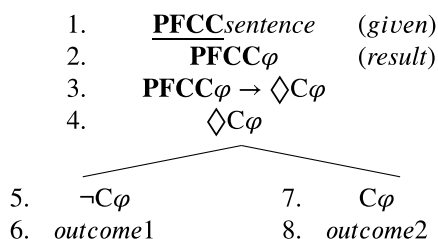
We thus, in a first step, assume **PFCC** and **PFCI** each to be functions that map sentences to *safely* prima facie conceivable sets of propositions or possible objects. This sounds really technical, hence it is opportune to ‘look at this in action’; we assume a dead-end act of prima facie imagination (**PFCI***sentence*) and observe the implications of what we have postulated:

²⁰The tradition underlying our approach can easily be grasped given the example of counterpossibles; “the Lewis-Stalnaker semantics has it that, if there are no A-worlds, $A \Box \rightarrow B$ comes out automatically true. The conditional with the same antecedent and opposite consequent, $A \Box \rightarrow \neg B$, comes out true, too, for the same reason. In general, all counterpossibles are vacuously true. The standard treatment of counterfactuals implies vacuism about counterpossibles.” (Berto & Jago, 2019, 267). Our approach consists in keeping the vacuism assumption and instead denying that our intuitions can refer to (conceptions related to alleged propositional contents of impossible antecedents of) counterpossibles. Cf. also (Goodman, 2004) on Stalnaker’s tradition.

²¹This is indeed a reference to Chomsky, who used a very similar, by now famous sentence (*colourless green ideas sleep furiously*) to demonstrate what he holds to be the semantic meaningfulness of some grammatically well-formed sentences (cf. Chomsky, 1957, 15). This connects with our inquiry insofar as dead-end prima facie conception might *correspond* to meaningfulness but also, at times, correlate with *some* different, and somehow related act of safe prima facie conception.



Here we show the branching of the proposed notion of dead-end prima facie conception:



So it seems that what in the literature is often just called prima facie conception actually could well collapse into the following subtypes of prima facie conception, which arguably together might constitute the insides of the mystery box ($\boxplus$). Together, they seem to provide an account of (the first layer) of the logical substructure of *general prima facie* (**GPF**):

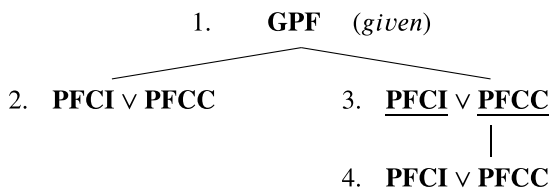
- (1) $\mathbf{PFCI}\varphi \rightarrow \Diamond \mathbf{I}\varphi$ (safe prima facie imagination implies ideal imaginability)
- (2) $\mathbf{PFCC}\varphi \rightarrow \Diamond \mathbf{C}\varphi$ (safe prima facie conception implies ideal conceivability)
- (3) $\underline{\mathbf{PFCC}}_{\text{sentence}} \rightarrow \mathbf{PFCC}\varphi$ (dead-end prima facie conception implies *safe* prima facie conception)
- (4) $\underline{\mathbf{PFCI}}_{\text{sentence}} \rightarrow \mathbf{PFCI}\varphi$ (dead-end prima facie imagination implies *safe* prima facie imagination)

We are aware of the fact that this above account does not say anything yet about how the mapping of sentences to *apparently* unrelated, ‘alternative’ propositions in the case of both dead-end prima facie types could look like. We will come to this later. General prima facie, what we tried to explore employing the mystery box metaphor ($\boxplus$), can thus possibly be understood as follows:

$$\mathbf{GPF} \rightarrow \neg[(\Diamond \mathbf{PFCI} \vee \Diamond \mathbf{PFCC}) \wedge (\Diamond \mathbf{PFCI} \vee \Diamond \mathbf{PFCC})]$$

This says that any act of general prima facie conception, if it leads to something, must either lead to (A) a *safe act* of prima facie conception/imagination of a *proposition* or to (B) a *dead-end act* of prima facie conception/imagination involving, in

some way, a *sentence*. And if B is the case, this means that actually something else is possibly *safely* conceived of *prima facie*:²²



So far we employed the metaphor of the mystery box when we used the naive notion of *prima facie* conception (PFC $\boxplus$). We will now *try* to get rid of this metaphor and of naive *prima facie* by first attempting to provide a general account of *prima facie* conception relying on the idea of it being a general function that takes in a sentence (or a set of sentences) *S* – just like we assumed it to be the case for the two obscure subtypes of *prima facie* conception. This means, we assume general *prima facie* conception to be a function that takes in as a variable a well formed *linguistic entity*, a *piece of information* *S* (a physical entity made of written words or spoken sounds that is processed in the brain).

This means we assume that brain states are formed as a result of the reception of a syntactically well-formed sentence *S* by a competent speaker. We call this the *linguistic parsing* of the linguistic entity *S*. Since the well-formedness of a *sentence* does not imply that the sentence expresses a conceivable *proposition* (for instance, in the case of the well-formed sentence ‘colourless ideas sleep squared circles’), we assume that a sentence *S* either expresses a proposition (namely, precisely when, upon being parsed by a competent speaker, it straightforwardly leads to a mental state corresponding to a conceivable proposition) or it does not. We therefore understand linguistic parsing as a *cognitive process* that is not necessarily resulting in propositional understanding (see below).

Let us be as clear as possible here, to avoid confusion; According to this chosen wording, one does never *linguistically parse* propositions. One parses a linguistic entity that can *express* a proposition. In case it does express a proposition, one can eventually imagine/conceive of the conveyed proposition; Following García-Carpintero and Palmira (2022), one can thus say that, *in the successful case of linguistic parsing*, the brain states involved during linguistic parsing – as well as the sentence – both ‘served as vehicles’ for the conveyed propositional content. In this way, the mind ultimately can grasp the proposition, by imagining/conceiving the content that the propositional vehicles ‘carried’.

²²At this point, one might consider whether *prima facie* undecidability could not also be introduced following (Yablo, 1993) as an additional, separate subtype of *prima facie* conception/imagination. It might be reasonable to assume that there are matters that are neither *ideally* conceivable nor contradictory (one could introduce a new, third type of *prima facie* for this purpose, here depicted by the *overline*): $\overline{\text{PFC}}x \rightarrow (\neg\Diamond\text{PFCI}x \wedge \neg\Diamond\text{PFCC}x \wedge \neg\Diamond\text{PFCI}\neg x \wedge \neg\Diamond\text{PFCC}\neg x)$. However, we are not further inquiring into this idea here.

There are different approaches available in the literature on how such parsing functions that map a sentence (or a set of sentences) S to a proposition P might look like.²³ We are, however, not primarily interested in the details of the biological/psychological side of this question here, i.e. in how the respective mapping can be modelled, say, neurologically, but in the epistemological side, i.e. in how a mapping could be possible at all in cases of merely apparent conceptions (dead-end *prima facie* conception/imagination).

Given our possibilist dogma concerning absolute conceivability, in a first step we propose to think of general *prima facie* in terms of the following process (1–4):

1. **Parsing:** Receiving a sentence S , a cognitive agent has new sensory data; the data must be *linguistically parsed*. We take the involved brain states (the brain states that constitute *linguistic parsing*) to be different *in kind* from acts of *conception/imagination* and therefore from understanding; in opposition to the latter, we assume that the respective brain states
 - (a) are potential vehicles for propositional content,
 - (b) but can also fail to be vehicles for propositional content.²⁴
2. **The parsing-test**²⁵: This test ‘asks’ the following: Does the respective parsing from step 1 result in a *propositional* event in the mind of the recipient of S ? Does the respective mental event furthermore involve visual/sensory representations?
3. **IF** the answer to the first part of the parsing-test is positive, then S is taken to encode the respective proposition φ . If, furthermore, the mental event involves visual representation, φ is ‘passed on’ to **PFCI**: **PFCI** φ . If the answer to the first part of the parsing-test is positive, but the parsing did not involve visual representations, φ is passed on only to **PFCC**: **PFCC** φ . In this way, it is assured that a sentence S *can*²⁶ straightforwardly guide to conception and imagination and therewith possibly result in an *understanding* of the corresponding proposition (or in an adequate conception/imagination of a possible object).
4. **ELSE**: S is ‘passed on’ to the dead-end *conception/imagination*-types: **PFCIS** or **PFCCS** that in turn result in **PFCI** φ or **PFCC** φ . This says nothing else than what we assumed above, namely that **PFCI** and **PFCC** are substructural

²³ Compare for instance the ideas on the history of the research in the field of *semantic parsing* in: Zhang et al. (2025).

²⁴ Sentences and brain states resulting from parsing of whole sentences can fail to carry and convey propositional content – they can fail to serve as ‘vehicles for meaning’. Just think, again, of the reception of the sentence “I am sitting at the library and I am reading and I am not reading”. The sentence as a whole fails to convey something that can have a truth value with respect to any logically possible world and the respective brain states thus fail to serve as vehicles for propositional content – at least at a face-value reading, i.e. ‘on the top-structure’.

²⁵ It is very important to emphasise that by using the term ‘test’, we do not wish to suggest that this is a process consciously carried out by the agents, a test that would provide them with any introspective clarity regarding the relationship between sentence, parsing and proposition. Here, we are describing the agent from the epistemological bird’s eye perspective, not their phenomenology.

²⁶ Remember the branching of *prima facie* conception shown above.

functions that somehow map a sentence to a proposition.²⁷ (This is thus still a problematic interim assumption).

Now a huge problem arises that we had postponed so far: If we just assume that the respective sentence expresses what could be the *content* of a dead-end, i.e. apparent act of conception/imagination (4), there seems to be no way of finding out how φ relates to S . For there seems to be no way to find out what an agent *prima facie* conceives of in case the agent undergoes a process of linguistic parsing that does *not* lead to propositional imagination/conception. For if it is assumed that some linguistic parsings do not end up in propositional understanding (2–4), and if only non-contradictory sets of propositions (and possible objects) can be conceived of/imagined, there seems to be no way to even write a paper about the nature of the relation between φ and S : For one would have to be able to isolate at least one case for which a statement about the content of a dead-end conception is possible, so that one could say something about what connects dead-end conceptions with *safe* *prima facie* conceptions more generally. Under our premises, however, any attempt to describe the content (allegedly) encoded by a sentence that leads to a dead-end conception would amount to an attempt to describe the *propositional content* of something that *cannot express a proposition* – that sounds unpromising.

It thus appears that, given our possibilist dogma concerning conceivability and imaginability (our path), this is the point where our inquiry into the substructure of *prima facie* conception and imagination must have an end; In the following subsection, however, we will argue that there *is* a way to further explore the deeper structure of the dead-end branches of general *prima facie* (without accepting, to Wittgenstein's horror,²⁸ that contradictions per se could somehow have propositional content).

4 The deep substructure of *prima facie* conception

We now attempt to inquire further into what we have called the *dead-end* types of general *prima facie* (GPF), i.e. those types of general *prima facie* that are still obscure:

- $\text{PFCC}_{\text{sentence}} \rightarrow \text{PFCC}\varphi$ (dead-end *prima facie* conception implies *safe* *prima facie* conception of *something*),
- $\text{PFCI}_{\text{sentence}} \rightarrow \text{PFCI}\varphi$ (dead-end *prima facie* imagination implies *safe* *prima facie* imagination of *something*).

So far, we have proposed to interpret these two subtypes of general *prima facie* merely in terms of functions from a sentence (or a set of sentences) S to a proposition φ (and so far we have not proposed any account of the workings of said function). What we have done so far is only to hypothesise that those intuitions that we *believe* to be related to a proposition arising from a dead-end act of *prima facie* conception

²⁷ Given our assumptions: $\text{PFCI}_{\text{sentence}} \rightarrow \text{PFCI}\varphi$ and $\text{PFCC}_{\text{sentence}} \rightarrow \text{PFCC}\varphi$.

²⁸ Essentially we accept that: “[...] contradictions [...] do not represent any possible situations.” (Wittgenstein 1961, 4.462).

of S are actually concomitant to a different, *safe* act of *prima facie* conception. – I.e. are actually, in a way, *about* a different proposition φ , different from *whatever impossibilists believe to be expressed by the respective S* .

We propose the following, in a certain sense Kantian, way to illuminate the dead-end subtypes; we believe one can ask about *the condition for the possibility* for why certain mental events cannot function as vehicles for propositional content: In other words, one can ask about the conditions of possibility for why certain parsing processes *fail* to be ‘vehicles’ for propositional content, given that one has a concept of what makes processes of linguistic parsing successful in conveying a propositional content from sentences to the mind. We call this *the transcendental argument*²⁹ *for the deep logical substructure of prima facie conception*:

- (i) The brain states corresponding to the linguistic parsing of a linguistic entity S *can* be vehicles for propositional content (a), but can also fail to be vehicles for propositional content (b).
- (ii) The parsing of a linguistic entity S either results in a *safe* act of *prima facie* conception/imagination or in a *dead-end* act of *prima facie* conception/imagination.
- (iii) Those *prima facie* acts that are *safe* guides to conceivability/imaginability are safe if and only if the brain states corresponding to the linguistic parsing *convey* a proposition, if they act as a vehicle for propositional content.
- (iv) Only non-contradictory sets of propositions are conceivable: $\Diamond C \Delta \longleftrightarrow (\Delta \not\vdash \perp)$.
- (v) A condition for the possibility of any $\times$ to qualify as a vehicle for propositional content is that $\times$ is non-contradictory. (This is the transcendental premise that gives the argument its name. We will soon define what ‘contradictory’ shall mean in this context).

ic A mental event that is non-contradictory *qualifies* as vehicle for propositional content and thus is, ideally speaking, a *safe* guide to conceivability/imaginability.

C A dead-end case of *prima facie* conception/imagination must either involve a contradiction or fail to serve as a vehicle for propositional content *in some other way*.

i and ii together correspond to the idea we developed above on general *prima facie*. iv is the definition of our dogmatic, possibilist path. The intermediate conclusion ic (from i-v) is the positive framing of the conclusion C. Paraphrased, the argument states that if something is not a mental state that involves an attitude towards a proposition (yet), but a process of linguistic parsing that *can function as a vehicle for propositional content*, the corresponding brain states that constitute the parsing must be non-contradictory. And therefore, if the opposite is the case, it must be so either because the respective brain states are contradictory or possibly because of some other reason.

But what could it mean for the brain states corresponding to linguistic parsing to be contradictory? How does this help us to investigate the logical substructure of dead-end types of *prima facie* conception? We believe the implications of the conclusion are insightful for the following reasons:

²⁹ Concerning transcendental arguments in contemporary philosophy cf.: Gram (1977), Kuusela (2008).

Let us recur to the mystery box we thought we had already got rid of; given that a *prima facie* act, if it does *not* ideally lead to propositional understanding, must be either contradictory (or something else must be held responsible for it *not* qualifying as a vehicle for propositional content), a dead-end *prima facie* act of imagination/conception can possibly be broken down as follows.

Metaphorically speaking, we have by now opened the big mystery box $\boxplus$; i.e. we have divided $\boxplus$ into the two transparent, safe types of *prima facie* conception (**PFCC** and **PFCI**) and two problematic, still in a way ‘obscure’ types of *prima facie* conception (**PFCC** and **PFCI**). One could represent our current insight by colouring in the mystery box, highlighting in black the two possible obscure cases of *prima facie* conception: $\blacksquare$. Although we have defined the obscure types in that we have assumed that they ‘throw back’ to a safe type, they are still obscure insofar as we have not specified how their ‘throwing back’ to a corresponding safe type of *prima facie* conception/imagination takes place.

The core idea of this section now is that, relying on the line of reasoning presented in *the transcendental argument for the deep logical substructure of prima facie conception*, one can ask about the conditions of possibility for why a dead-end case of *prima facie* conception might be obscure, treating the unsafe types like two substructural, nested mystery boxes $\blacksquare$ and $\blacksquare$ in $\blacksquare$:

For if one assumes that a dead-end *prima facie* act can arise from the combination of safe *prima facie* acts, i.e. from two *prima facie* acts of conception/imagination that *per se* each convey an *imaginable/conceivable proposition or object*, the following possible *causes* for the contradictory nature of **PFCC** and **PFCI** seem to arise:

- (I) two³⁰ safe, each *per se* imaginable, yet taken together contradictory acts of *prima facie imagination* resulting from the parsing of *S*.
- (II) two safe, each *per se* conceivable, yet taken together contradictory acts of *prima facie conception* resulting from the parsing of *S*.
- (III) a safe act of *prima facie imagination* that contradicts a safe act of *prima facie conception* resulting from the parsing of *S*.
- (IV) something else (for example, a consequence of humans being poorly built machines prone to error, or, you probably expected it, we could take this possible cause to be a *sub-sub-structural* mystery box, of a substructural mystery box, of the mystery box: as a $\blacksquare$ in $\blacksquare$ in $\blacksquare$).

Although examples are somewhat out of place here, as we do not wish to contribute to our understanding of how the brain *actually* works in these cases, a brief example follows to illustrate this idea, which should therefore be taken with a pinch of salt:

Suppose an agent receives the sentence ‘I am sitting in the library and I am reading and also not reading’.

The assumption we have just made is that the resulting totality of brain states generated by parsing in a competent speaker cannot function as a vehicle for propositional content in such cases. However, we assumed that this does not preclude, say,

³⁰ Of course, more than two conflicting states are possible, but this is covered in point IV. See further down in the text.

the first part of the sentence, ‘I am sitting at the library and I am reading’, from conveying a substructural proposition, which is entirely imaginable. Since this is a case in which two scenarios, each of which is in itself sensorially imaginable (‘I am sitting in the library and I am reading’ and ‘I am sitting in the library and I am not reading’), are in contradiction with one another (there is no logically possible world that could make this scenario true), this would be a case of the category I.

It thus seems that even without abandoning the possibilist dogma concerning conceivability, it is possible to distinguish between different *possible causes* for different dead-end prima facie conceptions/imaginings. And importantly, *some* of those possible causes must provide *safe*, propositional grounding (I-III). For it seems that either I-III must be the case or that, well, *something else* (IV) must be responsible for dead-end prima facie conceptions/imaginings. The two obscure kinds of prima facie conception (**PFCC**, represented by ■, and **PFCI**, represented by ■’) can thus, so we propose, be partly illuminated by non-dogmatic (we still leave IV open), transcendental reasoning:

If the parsing of a sentence (or a set of sentences) S does not convey a proposition, this does not exclude that the respective sentence/set of sentences can nevertheless be a vehicle for *substructural*, safely prima facie imaginable/conceivable propositional content. The cases I-III seem to reflect the idea of a possible *propositional substructure* of sentences that do not convey a proposition. Remember the principles of safe prima facie conception distribution and prima facie imagination distribution?

$$[(\varphi \wedge \psi) \not\vdash \perp] \rightarrow [(\Diamond \mathbf{PFCC} \varphi \wedge \Diamond \mathbf{PFCC} \psi) \longleftrightarrow \Diamond \mathbf{PFCC}(\varphi \wedge \psi)],$$

safe prima facie conception distribution,

$$[(\varphi \wedge \psi) \not\vdash \perp] \rightarrow [(\Diamond \mathbf{PFCI} \varphi \wedge \Diamond \mathbf{PFCI} \psi) \longleftrightarrow \Diamond \mathbf{PFCI}(\varphi \wedge \psi)],$$

safe prima facie imagination distribution.

We note that these principles do not exclude what we have just proposed concerning an understanding of the propositional substructure of contradictions; for prima facie imaginability/conceivability of contradictory matters is indeed excluded by these principles, *but it is not excluded* that two respective (substructural) conceivable/imaginable propositions can be imagined/conceived of ‘separately’ from each other.

5 The substructure of prima facie and intuitions

Some of what has been discussed here touches on a long-standing dispute in the philosophical tradition; for what we have observed is that in complex matters, we cannot rely on prima facie conception, since we do not have an introspective method to understand exactly when we are dealing with a possible conception/imagination that is mappable to the linguistic entity *we believe to grasp*, and when it is not. It seems we cannot make prima facie conception ‘transparent’. Just think of the example of the unsuccessful mathematicians given by Priest that was discussed above.

This lack of a method to render *prima facie* conception transparent is called the ‘Leibnizian objection’ because, in his *Meditations on Knowledge, Truth, and Ideas*, Leibniz raised this objection to Descartes’s notion of conception (Oliveri, 2021).

Note, however, that it was not our aim to make acts of *prima facie* conception transparent in the Cartesian sense of being evident and, therefore, reliable. Mostly the aim of the paper was to show that there is a possibilist, dogmatic way of understanding *prima facie* conception in cases that imply a logical contradiction. It was also not our ambition to provide arguments *against* the conceivability of impossibilities. The result is therefore mainly that a purely possibilist modal epistemology seems anything but transparent; for it may happen more often than we think that *what we believe to be expressed by a sentence* takes us astray and that we have no good way of finding out (by, say, introspection).

Nevertheless, we want to say a few words about the relation of the findings to intuitions (again, not understanding intuitions in the sense of ‘immediate’, ‘transparent’ conceptions/imaginings, as is done in some places in philosophy, but in the sense of phenomenal correlates to acts of imagination/conception that often lead us astray).

One can have the intuition that *p* even given the fact that *p* is false (Climenhaga, 2017) and there seems to be a “[...] lack [of] independent justification for the belief that intuitions are reliable” (Pust (2014), ch. 3.2).

Nevertheless, intuitions often have a certain argumentative ‘weight’ in many debates – philosophers, scientists (and most individuals of the species *Homo Sapiens*) *de facto* rely on intuitions. And sometimes, intuitions apparently can have so much ‘weight’, that we are prepared to accept *intuitions as evidence* that is *believed* to justify far-reaching adjustments with regard to philosophical methods, theories and beliefs. It is obscure, however, how we can reliably distinguish between those intuitions we should rely on, and those we should not – or if we can truly discriminate at all (Climenhaga, 2017).

Sometimes intuitions seem to be of help, when it comes to assessing the plausibility of a *prima-facie* conception (for instance when mathematicians have an intuition about a certain proof system that is well suited for a specific kind of problem), and sometimes they lead us astray. It seems that what we have discussed in this paper might potentially help to develop a (logically conservative, i.e. normal) framework that discriminates helpful³¹ intuitions from those that are not.

For not only could it be argued that the present approach might pave the way to answering the question of what our intuitions sometimes are *about* when we are ‘on the edge of conception’ (they might be about safe *prima facie* acts of conception).

What we have called *safe substructural prima facie* conceptions, which could be taken to underlie *dead-end* *prima facie* acts, provide a propositional grounding that can be more or less visual in nature (PFCI and PFCC). Also, the substructural possibilities I–III (cf. above) obtained via the transcendental argument and the combinatorics of these two *prima facie* forms seem to offer something like possible criteria for the reliability of an intuition; for it could be assumed that certain combinations (of substructural *prima facie* imagination and conception) have different *effects* on our

³¹ Helpful in the sense of ‘not leading us astray’, as can happen in the case of a ‘bad’ approach in mathematics in view of a conjecture *falsely* believed to be a theorem.

reasoning than others and could thus *guide* us better than other acts of *prima facie* conception.

It is important to acknowledge that we have not said much here about exactly how one might, through *introspection*, arrive at the conclusion that one particular intuition related to a conception/imagination is more reliable than another. It has merely been shown that, from a bird's-eye view of epistemology, a distinction can be drawn that may only be recognised by the perfect, ideally informed neuroscientist. Ultimately, therefore, it remains open how much we can draw from these bird's-eye-view findings for the functioning of intuitions.

Nevertheless, at least from the bird's eye view, the propositional nature of safe *prima facie* subtypes might make it possible to explain what the content of intuitions is, *if it should not be* the (alleged) content of dead-end *prima facie* conceptions: If one assumes, as is common, that the worlds in *W* are sorted in terms of their similarity, an assignment of the corresponding sentences to a *possible* world could be realised via the assignment of substructural safe *prima facie* conceptions; consider, for example, the following two counterpossibles as clarificatory examples:

- (1) If Wittgenstein's non-existent friend squared the circle in his private language, we should rethink our maths.
- (2) If the best mathematician alive would square the circle, we should rethink our maths.

It might, for instance, be possible to show that the substructural safe *prima facie* conceptions arising from the parsing of the antecedent in the second case (taken to be a linguistic entity at this stage) 'picks out' a closer possible world than the safe, substructural *prima facie* conceptions that might underlie the dead-end conception caused by the parsing of sentence one. This would possibly explain the different intuitions concerning the truth of 1 and 2 that philosophers report to have. And it wouldn't change the fact that we should perhaps continue to treat counterpossibles the same way, because it doesn't affect what we sometimes *believe* their antecedents express. In other words: the here proposed theory could help to show that what we sometimes apparently *intuit* to be our conception of *counterpossibles* (conditionals with impossible antecedents) actually corresponds to the conception of a substructural conditional with a logically possible antecedent. According to our 'path', epistemically speaking, counterpossibles *do not exist*.

Moreover, it should be noted that the possibility of recursive nesting of substructural *prima facie* acts (touched on in IV) arguably offers, logically speaking, the possibility of making infinitely fine-grained distinctions regarding the substructure of *prima facie* conceptions. This could, therefore, potentially help to make infinitely fine-grained distinctions when it comes to intuitions (again, from a bird's eye perspective of epistemology). This, too, demands its own discussion.

It seems that, if there were any truth in what has been said, this would affect those areas of philosophy that have to do with uncertainty, metaphysical disagreement, intuitions, provability, proof theory, semantic parsing, modelling of finite agents or that consist in the exploration of these very areas.

6 Results

It has been shown that, on the basis of a purely possibilist view (exclusively based on logically possible worlds), what Chalmers calls *prima facie* conception can likely be broken down as follows: in some cases, when we (naively speaking) conceive *prima facie*, safe acts of *prima facie* conception can be assumed and in other cases not. Initially relying on the metaphor of the mystery box (田), we have shown that (what we believe to be) the logical substructure of *prima facie* conception can be explored much further than initially assumed. To this end, we have endeavoured a transcendental argument that asks about the conditions of the possibility of safe and dead-end *prima facie* conception/imagination. That is, the conditions of possibility of those *prima facie* acts that, ideally speaking, necessarily throw us back to logical possibility and those that, ideally speaking, do never guide to logical possibility.

We proposed that dead-end cases of *prima facie* conception, those acts of *prima facie* conception that do never lead to possibility and thus have no possible content, could, so we suggested, be analysed in terms of a function that maps sentences to safe *prima facie* acts of conception/imagination. This seems to pave the way to what we called the (possibly recursive) logical substructure of *prima facie* conception.

If one is furthermore willing to accept our distinction between *sensorial* and *conceptual* *prima facie* acts (and possibly even *undecidable* *prima facie* types in the spirit of Yablo), this might provide a powerful approach that helps to elucidate (at least from the bird's eye view of epistemology) the as yet rather obscure relation between what we intuit conceive, what our intuitions actually refer to and what is ideally conceivable. We called the non-naive view of *prima facie* conception, the view that takes into account the substructure, general *prima facie* conception.

Instead of adapting formal semantics to our intuitions, we propose to inquire further 'into' 田. For it seems that even if one may never get fully rid of 田, there may still be a great deal of meaningful things to say about what we arguably conceive of when we (erroneously) intuit to conceive of impossibilities.³²

Acknowledgment I thank Luis Rosa for being a possibilist. I thank Greg Restall for his formidable ability to listen carefully and for his ideas. I thank Cristina Caroli Costantini for her mathematical intuitions and her kindness. I thank Jasper Lohmar, Nicola Thomaschewski, and all participants of the Brownbag Colloquium in Cologne for their helpful comments on earlier drafts of this paper. I thank Filippo Riscica Lizzio and Matteo Giannone for a wonderful time in St Andrews. And, as always, I thank the METRONOM association. Last but not least, I would like to express my sincere thanks to the anonymous referees of this journal.

Funding Open Access funding enabled and organized by Projekt DEAL.

Data availability Not Applicable.

³² We hope and argue it is not conceivable that the early Wittgenstein would turn and would also not turn over in his grave when confronted with what has been said.

Declarations

Conflict of interest No potential conflict of interest was reported by the author. The author has seen and agreed with the contents of the manuscript, and there is no financial interest to report. We certify that the submission is original work and is not under review at any other publication.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

- Antognazza, M. R., & Segala, M. (2023). Intuition in the history of philosophy (what's in it for philosophers today?). *British Journal for the History of Philosophy*, 31(4), 574–578. <https://doi.org/10.1080/09608788.2023.2186833>
- Arana, A. (2016). Imagination in mathematics. In A. Kind (Ed.), *The routledge handbook of the philosophy of imagination* (pp. 463–477). Routledge.
- Berto, F. (2014). V-On conceiving the inconsistent. *Proceedings of the Aristotelian Society (Hardback)*, 114(1pt1), 103–121. <https://doi.org/10.1111/j.1467-9264.2014.00366.x>
- Berto, F. (2017). Impossible worlds and the logic of imagination. *Erkenntnis*, 82(6), 1277–1297. <https://doi.org/10.1007/s10670-017-9875-5>
- Berto, F., French, R., Priest, G., & Ripley, D. (2018). Williamson on counterpossibles. *Journal of Philosophical Logic*, 47(4), 693–713. <https://doi.org/10.1007/s10992-017-9446-x>
- Berto, F., Hawke, P., & Özgün, A. (2022). *Topics of thought. The logic of knowledge, belief, imagination*. Oxford University Press.
- Berto, F., & Jago, M. (2019). *Impossible worlds*. Oxford University Press. <https://doi.org/10.1093/oso/9780198812791.001.0001>
- Brown, J. R. (1997). Proofs and pictures. *British Journal for the Philosophy of Science*, 48(2), 161–180. <https://doi.org/10.1093/bjps/48.2.161>
- Chalmers, D. J. (2002). Does conceivability entail possibility? In T. S. Hawthorne, & G. John (Eds.), *Conceivability and possibility* (pp. 145–200). Oxford University Press. <https://doi.org/10.1093/oso/9780198250890.003.0004>
- Chomsky, N. (1957). *Syntactic structures*. De Gruyter Mouton. <https://doi.org/10.1515/9783112316009>
- Climenhaga, N. (2017). Intuitions are used as evidence in philosophy. *Mind*, 127(505), 69–104. <https://doi.org/10.1093/mind/fzw032>
- Costa-Leite, A (2010) Logical properties of imagination. *Abstracta*. 6(1), 103–116. <https://doi.org/10.2433/8/abs-2010.258>
- Descartes, R. (1996) Sixth meditation: The existence of material things, and the real distinction between mind and body. *Cambridge Texts in the History of Philosophy*, 50–62. Cambridge University Press, Cambridge, Introduction by Williams, Bernard. 50 ff. 1996
- Duc, H. N. (1997). Reasoning about rational, but not logically omniscient, agents. *Journal of Logic and Computation*, 7(5), 633–648. <https://doi.org/10.1093/logcom/7.5.633>
- García-Carpintero, M., & Palmira, M. (2022). What do propositions explain? Inflationary vs. deflationary perspectives and the case of singular propositions. *Synthese*, 200(2), 107. <https://doi.org/10.1007/s11229-022-03467-7>
- Gendler, T. S., & Hawthorne, J. (2002). Introduction: Conceivability and possibility. In T. S. Gendler, & J. Hawthorne (Eds.), *Conceivability and possibility* (pp. 1–70). Oxford University Press.

- Goodman, J. (2004). An extended lewis/stalnaker semantics and the new problem of counterpossibles. *Philosophical Papers*, 33(1), 35–66. <https://doi.org/10.1080/05568640409485135>
- Gram, M. S. (1977). Must we revisit transcendental arguments? *Philosophical studies. Philosophical Studies*, 31(4), 235–248. <https://doi.org/10.1007/BF01855229>
- Hilbert, D., & Cohn-Vossen, S. (1952). *Geometry and the imagination*. Chelsea Publishing Company.
- Kind, A. (2023). Why we need imagination. In B. McLaughlin, & J. Cohen (Eds.), *Contemporary debates in philosophy of mind* (2nd ed., pp. 570–587). John Wiley and Sons, Ltd.
- Kuusela, O. (2008). Transcendental arguments and the problem of dogmatism. *International Journal of Philosophical Studies*, 16(1), 57–75. <https://doi.org/10.1080/09672550701667083>
- Oliveri, L. (2021). Conceivability errors and the role of imagination in symbolization. *JoLma*, 2(2), 293–312. <https://doi.org/10.30687/Jolma/2723-9640/2021/02/002>
- Priest, G. (2016). Thinking the impossible. *Philosophical Studies*, 173(10), 2649–2662. <https://doi.org/10.1007/s11098-016-0668-5>
- Pust, J. (2014 Fall). Intuition. In E. N. Zalta, & U. Nodelman (Eds.), *The stanford encyclopedia of philosophy, Fall 2014 edn. Metaphysics research Lab*. Stanford University, Stanford. <https://plato.stanford.edu/archives/fall2014/entries/intuition/>
- Salmon, N. (2024). From modality to millianism. *Noûs*, 59(4), 851–872. <https://doi.org/10.1111/nous.12536>
- Stuart, M. T. (2022). Thought experiments. In V. P. Glăveanu (Ed.), *The palgrave encyclopedia of the possible* (pp. 1644–1654). Springer. https://doi.org/10.1007/978-3-030-90913-0_59
- Tanaka, K. (2018). Logically impossible worlds. *The Australasian Journal of Logic*, 15(2), 489–497. <https://doi.org/10.26686/ajl.v15i2.4870>
- Wansing, H. (2017). Remarks on the logic of imagination. A step towards understanding doxastic control through imagination. *Synthese*, 194(8), 2843–2861. <https://doi.org/10.1007/s11229-015-0945-4>
- Williams, T. (2025). Anselm of Canterbury. In E. N. Zalta, & U. Nodelman (Eds.), *The stanford encyclopedia of philosophy*, (Summer edn). Metaphysics Research Lab, Stanford University, Stanford. <https://plato.stanford.edu/archives/sum2025/entries/anselm/>
- Williamson, T. (2007). *The philosophy of philosophy*. Blackwell Publishing.
- Wittgenstein, L. (1961). *Tractatus Logico-Philosophicus*. Routledge & Kegan Paul.
- Yablo, S. (1993). Is conceivability a guide to possibility? *Philosophy and Phenomenological Research*, 53(1), 1–42. <https://doi.org/10.2307/2108052>
- Zhang, X., Bouma, G., & Bos, J. (2025). Neural semantic parsing with extremely rich symbolic meaning representations. *Computational Linguistics*, 51(1), 235–274. https://doi.org/10.1162/coli_a_00542

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.